

ELECTRONIC FUNDAMENTALS

THEORY LESSON 30

OSCILLATORS

- 30-1. Basic Oscillator Theory
- 30-2. Tickler-Coil Oscillator
- 30-3. Hartley Oscillator
- 30-4. Colpitts Oscillator
- 30-5. Tuned-Plate Tuned-Grid
Oscillator
- 30-6. Crystal Oscillator
- 30-7. Miscellaneous Oscillators
- 30-8. Coupling



RCA INSTITUTES, INC.
A SERVICE OF RADIO CORPORATION OF AMERICA
HOME STUDY SCHOOL
350 West 4th Street, New York 14, N. Y.

Theory Lesson 30

INTRODUCTION

There are many cases in which it is necessary to produce an alternating voltage of a particular frequency. Although it is possible to use a mechanical generator to produce alternating voltage, the maximum frequency that can be generated is limited. Also, there are certain physical limitations upon the use of generators. Furthermore, it is not a simple matter to hold the frequency of a generator constant. It would be a more complex task to vary the frequency of the voltage output.

However, an oscillator, which is an electrical circuit that produces an alternating current of a desired frequency, can be used where the use of mechanical generators is not practicable.

In the lesson on detection, you learned that at the broadcasting station, a carrier wave is caused to vary in amplitude at an audio rate in order to modulate an r-f carrier so that AM (amplitude-modulated) waves can be transmitted. You know that the audio component of the AM wave is produced by voice or music sounds fed into a microphone, which, in turn is coupled to the speech-amplifier section of a transmitter.

But how is the carrier wave generated? It originates in an *oscillator*.

The oscillator is also used in the superheterodyne receiver. A block diagram of a superheterodyne receiver is shown in Fig. 30-1. Note that the block below the mixer stage is called the *local oscillator*. Like the oscillator in the transmitter, the local oscillator generates an r-f wave. The frequency of the r-f wave generated by the local oscillator differs from the incoming signal frequency by an amount that is just equal to the intermediate frequency.

In addition to its use in receivers and transmitters, the oscillator is used to generate a.c. in test equipment, television sets, radar, and many other devices using electronic circuits.

The common oscillators consist of an electron tube, a tuned circuit, a source of d-c power, and a feedback circuit. As mentioned in previous lessons, a circuit has feedback when a portion of the output is fed back to the input. If the feedback is opposite in phase to the input or signal voltage, it is known as degenerative feedback, and will cause a decrease in the input signal voltage. When the feedback

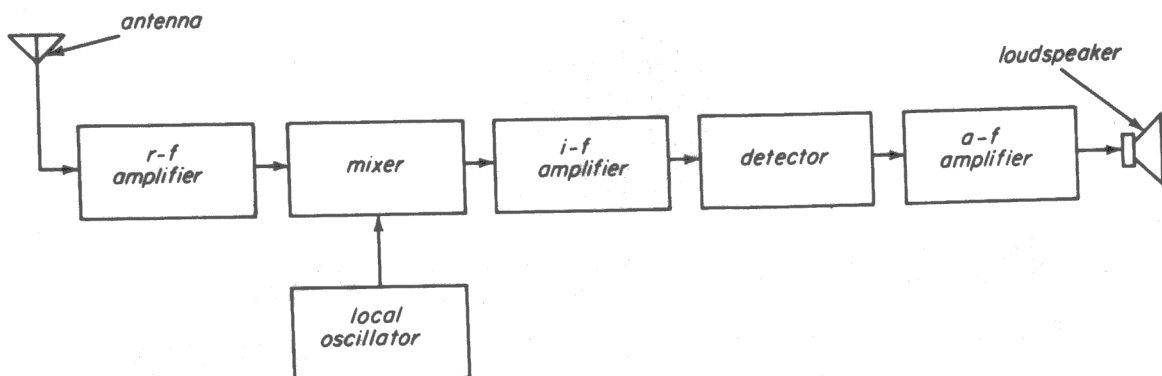
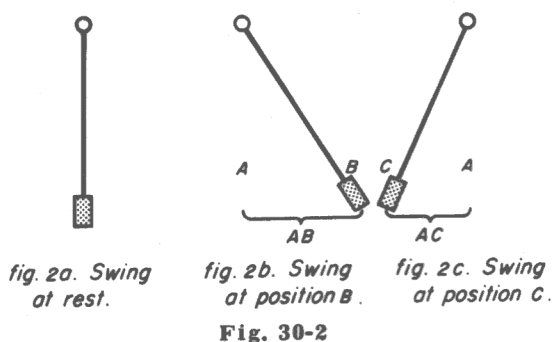


Fig. 30-1



voltage is in phase with the input, it will add to the input, and therefore increase it. This is regenerative feedback, the kind that we need to make an oscillator work. Before we discuss the theory of a basic oscillator circuit, let us start by pointing out the similarity between a mechanical and an electronic oscillator. We'll compare a swing with a tank circuit.

30-1. BASIC OSCILLATOR THEORY

A dictionary defines the word *oscillate* as follows: to swing to and fro. Does this mean that a child's swing can be called an oscillator, just because it goes back and forth? Exactly—a swing is a mechanical oscillator. Let us study the action of a swing. Figure 30-2a shows a swing at rest. If you push this swing, it will move forward from its resting position to a new position, as illustrated in Fig. 30-2b. At point B, the swing cannot go any farther. The swing starts back. On its return, it does not stop at its original resting position, but goes beyond to a new position shown in Fig. 30-2c. Inertia tends to keep moving objects in motion and in the same direction.

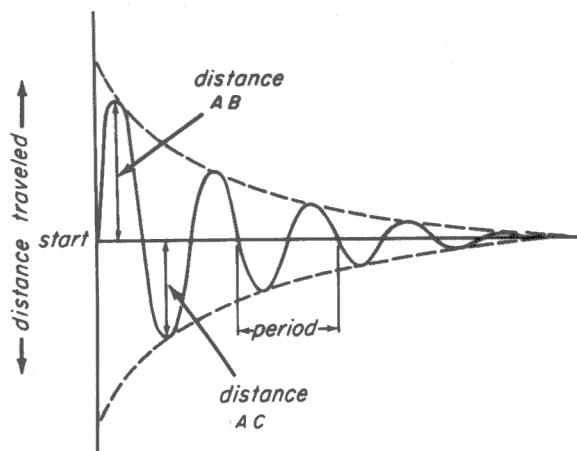


Fig. 30-3

Note in Fig. 30-2b and 2c that the distance between A and B is greater than the distance between A and C. This is true because a small portion of the energy transferred to the swing when it received a push has been lost due to friction and air resistance. The swing will continue to go to and fro—in other words, it will keep oscillating until all of its energy has been lost in friction and air resistance.

The graph in Fig. 30-3 plots the distances travelled by the swing. You can see that these distances grow smaller and smaller with each oscillation. But the length of time between oscillations remains the same. The first complete cycle takes the same amount of time as the last. The time for completing one full cycle is known as the *period* of oscillation. The number of cycles completed in one second is the *frequency*. Period and frequency are related by the following formula:

$$t = \frac{1}{f}$$

where:

t is the period

f is the frequency.

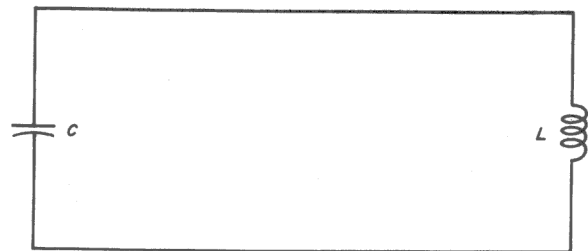
If you call the distance travelled by the swing from its resting position its amplitude, you can summarize as follows: the amplitude of oscillations decreases each successive period, but the frequency remains constant. This behavior of a swing can be compared to the action in a parallel-resonant circuit containing a coil and a capacitor. Such a circuit is known as a tank circuit.

The Tank Circuit. The tank circuit (Fig. 30-4a) is the simplest form of electrical oscillator. Let us move the switch S of Fig. 30-4b to position 1, thus connecting the source of d.c. across the capacitor. Instantly, electrons will be transferred from plate 1 of the capacitor to plate 2. In a fraction of a second, the capacitor will become charged. Plate 1 will be positive because it will have a deficiency of electrons; plate 2 will be negative because it will have an excess of electrons.

When the d-c source is disconnected, and the capacitor is connected across the coil, the capacitor will discharge and then be recharged to the opposite polarity, discharge a second time, and again be recharged to its original polarity. For convenience, we shall call the discharge period *first quarter*, the charging period *second quarter*, the next discharge period *third quarter*, and the final period *fourth quarter*.

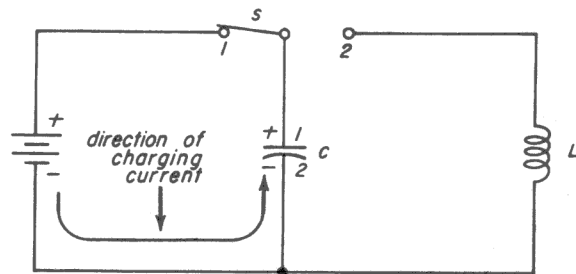
To start the action of the first quarter, we will move switch *S* to position 2, connecting the capacitor in parallel with the coil, (Fig. 30-4c). The capacitor will start to discharge through the coil. Plate 2 of the capacitor is negative: therefore, electrons will flow from plate 2 through the inductor to plate 1. At the instant electrons start to flow through the inductor, a magnetic field starts to build up around the coil. This expanding magnetic field cuts across the turns of the coil, and it induces a counter emf that is equal to the capacitor voltage, but opposite in direction. The induced emf opposes any change in the current flow through the circuit, and, therefore, it opposes the discharge current of the capacitor. The induced emf is greatest when current first starts to flow, because this is when the current is changing most rapidly. As the rate of change of current decreases due to the counter emf, the flux lines in the magnetic field change more slowly, and the counter emf decreases. The voltage of the capacitor is decreased, since some of its electrons have been transferred from one plate to the other, thus lowering the difference of potential across the capacitor. At any instant, the counter emf is equal to the capacitor voltage, but opposite in direction (polarity). When the current reaches its maximum value, its rate of change is zero. At this instant, both the capacitor voltage and the induced emf are zero. The capacitor is now completely discharged.

At the start of the first quarter, energy was stored in the electrostatic field of the capacitor. This energy was dependent upon the capacitance and the charging voltage. When the capacitor was discharged and its voltage dropped to zero, no energy remained in the electrostatic field. What became of this energy? A small portion of the energy



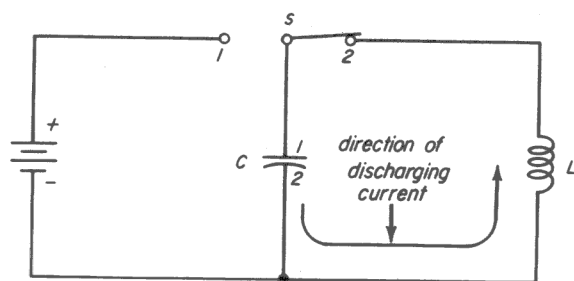
tank circuit

(a)



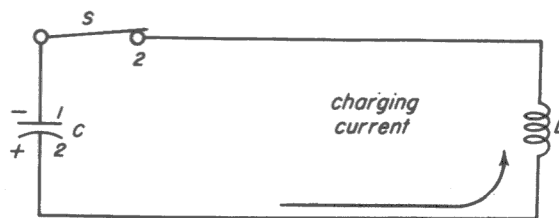
capacitor charging

(b)



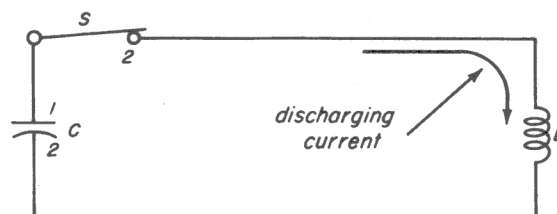
capacitor discharging

(c)



capacitor charging to opposite polarity

(d)



capacitor discharging, opposite direction

(e)

Fig. 30-4

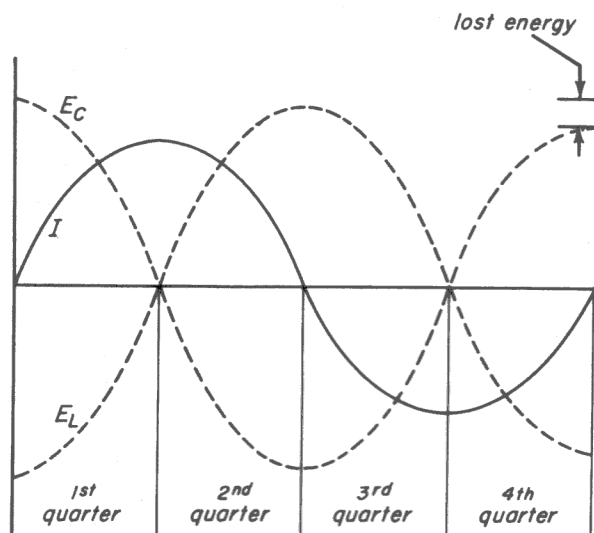


Fig. 30-5

was lost in the resistance of the circuit. The greater part of the energy was transferred to the magnetic field of the inductor. At the instant the discharge current reached its maximum value, the energy in the magnetic field was greatest. This is true because the amount of energy stored in a magnetic field depends on the inductance of the coil and the current flowing through the coil.

To review the action of the first quarter, the discharging capacitor caused a current to flow through the inductor. Most of the energy originally stored in the capacitor was transferred to the coil's magnetic field. At the instant this transfer was complete, current was maximum, capacitor voltage was zero, and induced emf was zero. The conditions during this interval are shown in Fig. 30-5. (E_C is the capacitor voltage; E_L is the counter emf; and I is the current).

In order for the magnetic field to be sustained, current must continue to flow, because the flow of current causes the magnetic field to build up. But current cannot flow, because it has reached a maximum value, and there is no difference of potential in the circuit to keep it flowing. Since there is nothing to sustain the magnetic field, it begins to collapse.

This is the start of the second quarter. The energy stored in the magnetic field starts to return to the electrostatic field of the capacitor. As the magnetic field collapses, its flux lines cut the turns of the

coil in a direction that is opposite to the expanding magnetic field of the first quarter. Once again, a counter emf is induced, but this time its direction is opposite to the induced emf of the first quarter. In other words the induced emf of the second quarter is of the same polarity as the capacitor voltage of the first quarter. Therefore, the effect of the induced emf will be to continue the flow of current through the coil in the same direction as in the first quarter. As the induced emf increases, more electrons are moved around to plate 1 of the capacitor (Fig. 30-4d). The capacitor is beginning to be recharged, but this time plate 1 is collecting an excess of electrons, and developing a negative charge. This negative charge opposes the flow of current, and causes it to decrease more rapidly each instant. The more quickly the current decreases, the greater will be the induced emf. At the instant the induced emf and the capacitor voltage reach their maximum values, the current is zero, and all the energy minus the losses is once again stored in the capacitor. Because of the energy lost in the circuit resistance, the charge on the capacitor is slightly less at the end of the second quarter than it was at the beginning of the first quarter. The polarity of the capacitor's charge is now opposite to its polarity in the first quarter (see Fig. 30-5).

The third quarter is a repetition of the first quarter, but in the opposite direction. The direction of current flow is illustrated in Fig. 30-4e. The energy is once again transferred to the magnetic field, and at the end of the third quarter, the charges on the capacitor will have equalized again, dropping the voltage across the capacitor to zero.

The fourth quarter repeats the second quarter, but the direction is opposite. At the end of this quarter, all circuit conditions will be exactly as they were at the beginning of the first quarter with this important exception—a certain portion of the energy has been lost (see Fig. 30-5). The series of events in the four quarters are considered a cycle, and the time for this cycle to take place is the period. After the first cycle, other cycles follow. The oscillations will

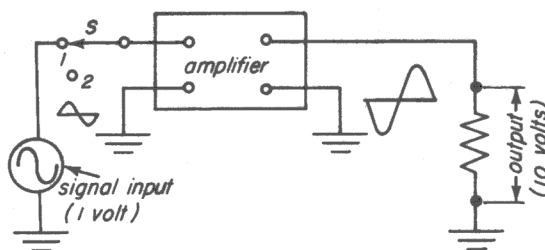
gradually die away, however, because of the continuous loss of energy in the circuit resistance. The waveforms of these oscillations are exactly like the graph in Fig. 30-3, which plots the distances travelled by the swing. The tendency of the swing to continue oscillating is similar to the tendency of the tank circuit to oscillate. In the case of the swing, the motion was sustained by mechanical inertia; in the tank circuit, it is the electrical inertia that continues the oscillations.

Circuit Losses. As the oscillator operates, the oscillations tend to die out, or become *damped*. The oscillations become damped because of losses in the circuit. Damped waves are not useful. So, the oscillators used to produce an alternating current that is of constant amplitude (not damped) are those in which circuit losses have been overcome.

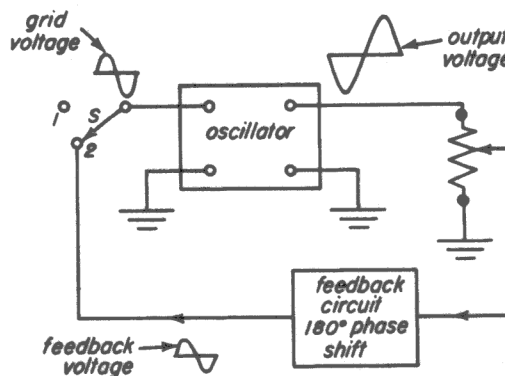
Let's consider the swing and how the losses in its oscillation might be overcome. If you push the swing each time it returns to a certain position in its cycle, the swing will go on oscillating indefinitely. Each time you push the swing, you are adding energy to it. Energy can be added to the tank circuit of an oscillator too. This extra energy will keep it oscillating.

This extra energy must be added to the circuit in such a direction that aids the flow of current. The source of this energy should be an emf that produces a current of the same frequency and direction as the oscillating current. Enough energy should be added to make up for all of the resistance losses in the circuit. In order to obtain enough energy to overcome circuit losses and keep the tank-circuit oscillating, the oscillator is designed to act as an amplifier. As a matter of fact, a feedback oscillator is an amplifier that supplies its own input, in such phase and amount as to cause and keep up oscillations.

Figure 30-6a is a block diagram of an amplifier. The input signal of one volt is amplified to ten volts in the output, and, at the same time, its phase is reversed. Compare the waveshapes of the input and the output voltages. Now let us disconnect the



(a)



(b)

Fig. 30-6

external input signal by moving switch S to position 2 and installing a feedback circuit (Fig. 30-6b). A portion of the output is shifted 180° and fed back to the input. Since the output of an amplifier is 180° out of phase with the input, the additional phase shift of 180° made the total shift equal to 180° plus 180°, or a total of 360°. This brings the feedback voltage into phase with the input voltage at the grid.

Requirements of Resonant Electron-Tube Oscillators. The requirements of an electron-tube oscillator are as follows. It must contain an oscillatory tank circuit, consisting of a capacitor and an inductor. The frequency of oscillation is determined by the values of L and C , according to the formula:

$$f = \frac{1}{2\pi\sqrt{LC}}$$

For small values of L and C , the frequency will be high, and vice versa. The circuit must have a source of d-c energy. Finally, there must be a feedback arrangement capable of supplying the right amount of energy in the proper direction. The tube must be able to amplify the input voltage.

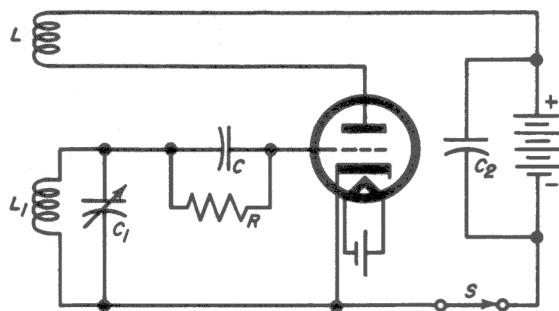


Fig. 30-7

30-2. TICKLER-COIL OSCILLATOR

One of the earliest electron-tube oscillator circuits is the *tickler-coil oscillator*, also named the *Armstrong circuit* after its inventor. Figure 30-7 is a schematic diagram of such a circuit. Coil L is the feedback coil. It is also called the *tickler coil*. L_1 and C_1 make up the oscillatory tank circuit, R is the grid leak resistor, and C is the grid capacitor. The B battery is connected to the plate through L , the feedback coil; therefore, both direct and alternating components of plate current flow through L . This arrangement is known as a *series-fed circuit*. The bypass capacitor C_2 offers a low reactance path for the alternating component of plate current, thus bypassing the battery for a.c. Once oscillations have been started in the tank circuit, they are amplified by the tube and appear in the output with a 180° phase reversal. The combination of L and L_1 form a transformer, and this introduces another 180° phase shift. Thus the feedback voltage induced in L_1 by L has gone through two phase shifts of 180° each, making a total of 360° . This brings the feedback voltage into phase with the grid voltage.

The instant switch S of Fig. 30-7 is closed, the plate voltage applied to the tube causes electrons to flow through the circuit. This current sets up an expanding magnetic field about coil L , and thus a voltage is induced in L_1 . This voltage will be in the proper direction to force electrons to move from grid side of capacitor C_1 to the cathode side. (The direction of the winding in the secondary circuit of the transformer is made such that the polarity of the induced secondary voltage is correct.) This causes the upper side of C_1 to develop a positive charge with respect to the bottom side. There is no charge on C . Thus, there is no bias on the

tube at the start; therefore, the positive voltage applied to the grid causes an increase in plate current. This builds up a greater magnetic field about L , and a larger emf is induced in L_1 . This increases the current in the tank circuit, making C_1 more positive. This process continues until plate current almost reaches the saturation point, where a further increase in grid voltage does not increase plate current. Then the magnetic field about L stops expanding. Since the flux lines in the magnetic field about L are not changing, a voltage is no longer induced in L_1 . Capacitor C_1 , which has been charged to the maximum positive voltage, begins to discharge through L_1 . This discharge current will be opposite in direction to the current that charged C_1 positively, because electrons from the negative side of C_1 will be flowing upward through L_1 and around to the grid side of C_1 . This makes the grid more negative than before, and plate current starts to decrease. The decreasing plate current causes the magnetic field about L to start collapsing, and the flux lines cut across the turns of L_1 as they move inward. Now the emf induced is in the proper direction to force electrons to flow from the cathode side of C_1 to the grid side. In other words, the emf that is induced in L_1 is in the proper direction to aid the discharge of C_1 . The grid is made more and more negative, until a point is reached where plate current is cut off.

At the instant of plate current cutoff, the upper side of the capacitor C_1 has a maximum negative charge, the field about L_1 has collapsed completely, and the voltage induced in L_1 has fallen to zero. Since there is no current flow, and nothing else to oppose the discharge of capacitor C_1 , it will discharge through L_1 in such a direction as to cause electrons to move from the grid side of C_1 to the cathode side. The grid will become more positive, plate current will start to flow once again, and the entire cycle that started when switch S was closed is repeated. One complete cycle takes place in a very short period of time, and it is repeated thousands of times in a second, at a frequency determined by the resonant formula:

$$f = \frac{1}{2\pi\sqrt{LC}}$$

Oscillations continue indefinitely if the energy fed back by coil L to the tank circuit is enough to overcome losses in the circuit. The amount of feedback depends on the mutual inductance between L and L_1 ; the closer the two coils are brought together, the greater will be the amount of feedback. As the coils are moved farther apart, the coupling becomes looser, and the amount of feedback decreases. There is a critical point in the coupling between the coils beyond which the coupling is too loose for oscillations to continue, because the feedback voltage is too small at that point. In adjusting the coupling in an Armstrong oscillator, care should be taken not to make it too loose, nor too tight.

Grid-Leak Bias. The previous discussion explains how the d-c power in the plate circuit starts a chain of events that finally develop an a.c. in the tank circuit. The a.c. has a frequency that is determined by the L and C of the tank circuit. Let us now examine the function of the grid-leak resistor R and grid capacitor C of Fig. 30-7. In order to obtain high efficiency in an oscillator, the tube is usually operated class C . This requires that the tube be biased beyond cutoff normally between $2\frac{1}{2}$ and 4 times the value of bias necessary for cutoff. The use of grid-leak resistor and capacitor is required, so that grid-leak bias can be developed. Fixed bias greater than the cut-off bias would not allow the oscillator to be self-starting, because the current would be cut off when the switch is thrown. When grid-leak bias is used, there is no initial bias on the tube when the plate voltage is applied.

The graph in Fig. 30-8 illustrates how grid-leak bias is developed. As the oscillations build up from a small value to their constant amplitude, the swing of grid voltage goes positive for a small portion of each cycle. When the grid is driven positive with respect to cathode, grid current flows and charges C , making the grid end of C negative. The schematic diagram for the tickler feedback oscillator has been redrawn in Fig. 30-9a. Notice that portions of diagram appear in heavy lines. As shown in b of the figure, the tank circuit can be

constitutes the grid bias for the oscillator. In a certain length of time, the grid bias reaches its proper operating value, and then remains constant. This length of time is a fraction of a second, and depends upon the time constant of R and C . The time constant of an RC circuit is equal to the product of R and C , and is expressed in seconds. Note that, if the tank voltage increases, the bias across C will also increase. A resistor is placed across C . The resistor is called the grid-leak resistor. The purpose of the grid-leak resistor is to allow the grid bias (across) C to decrease whenever the peak tank voltage becomes less than the d-c grid bias. It is clear from this discussion that the time constant of the grid resistor-capacitor combination is rather critical in the operation of an oscillator. If the values of R and C are chosen so that the time constant is too short, C will discharge through R during the negative portions of the grid voltage cycle, and will not give the bias a chance to develop to its proper operating value. On the other hand, the time constant should not be too long, because the grid bias could reach a value high enough to stop the flow of plate current, and thus stop the oscillations.

Look back to Fig. 30-8. The graph illustrates the flow of plate current in a class C oscillator. The tube is biased beyond cutoff; this means that current flows in the plate circuit for *less* than half a cycle of the oscillating tank circuit. The distortion of the plate-current waveform is not reproduced in the grid circuit because the tank circuit stores energy in the capacitor and coil on alternate half-cycles and generates an oscillating sine-wave current.

The action of the grid-leak circuit actually prevents the plate of the tube from drawing saturation current. The grid-leak regulates the action of the oscillator so that plate current reaches a constant value, sufficient to transfer the correct amount of energy to the grid circuit. This energy is just enough to keep the amplitude of oscillation large enough to continue plate current flow at the constant value mentioned above. There is no surplus energy to build up the amplitude of oscillations any further. Notice

in Fig. 30-8 that this constant value of plate current is reached before the grid bias reaches its proper operating point, because the tube operates beyond cutoff. The action of the tank circuit causes the amplitude of the oscillations to be large enough to drive the grid slightly positive for a small portion of each cycle. Because of this, the right amount of grid current flows to maintain a constant value of bias.

Oscillator Stability. A very important factor in the operation of an oscillator is the ability to maintain a constant frequency for long periods of time. It has been found that frequency stability is highest when oscillations occur at a frequency that is as close as possible to the resonant frequency:

$$f_r = \frac{1}{2\pi\sqrt{LC}}$$

The exact frequency of oscillations of an electron tube oscillator is equal to the resonant frequency *multiplied* by another term:

$$f_o = \left(1 + \frac{R}{R_p}\right) f_r$$

where:

R is the total effective series resistance of the tank circuit

R_p is the a-c plate resistance of the tube.

When the effective series resistance of the tank circuit is low, the term $\frac{R}{R_p}$ is so small that it can be disregarded. The Q of a tank circuit is equal to:

$$\frac{X_L}{R}$$

For high values of Q , R is correspondingly low. In considering *no-load* conditions, the tank circuit Q should be high, in the order of 80 to 150.

Loading. When a load is reflected into the tank circuit, (Fig. 30-10a), the Q is lowered, and the frequency of oscillations is shifted. The effect of the load is to reflect an impedance into the tank circuit.

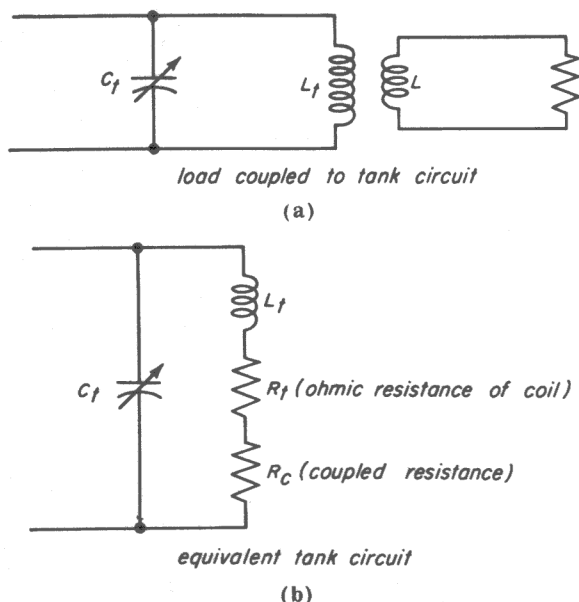


Fig. 30-10

The reflected or coupled impedance consists of a resistive component and a capacitive reactance. The reactive component of the reflected impedance is capacitive because the coupling coil L passes current that induces a counter emf into the tank circuit, and this counter emf is 180° out of phase with the voltage of the coupling coil. Since the coil L represents an inductive reactance, the coupled reactance into the tank circuit is 180° out of phase, and thus it is capacitive. Any consequent detuning may be overcome by retuning C_t .

The approximate amount of reflected or coupled resistance can be found from the formula:

$$R_c = \frac{(2\pi fM)^2}{R_s}$$

where:

R_c = resistance in ohms reflected into the primary.

R_s = load resistance in ohms across secondary.

M = mutual inductance in henrys.

f = frequency in cps.

This resistive component of the coupled impedance is in series with the ohmic resistance of the tank coil, (Fig. 30-10b). The circuit Q is now equal to $XL/(R_t + R_c)$. You can see that the Q of an oscillator is lowered when it delivers power to a load, be-

cause a resistance is added to the tank circuit. As an example, let us consider a circuit in which $X_L = 5,000$ ohms and $R_t = 50$ ohms. The Q of the circuit is $\frac{5000}{50}$ or 100, when a load is not connected. Suppose we couple a load to this circuit, and the coupled resistance is 350 ohms. The Q of the same circuit when loaded is $\frac{5000}{350 + 50}$ or 12.5. In other words, the effect of loading this oscillator circuit is to reduce the Q from a value of 100 to a value of 12.5.

The Q of an oscillator tank circuit when loaded is a measure of the ratio of power stored in the tank circuit to the power lost in the circuit resistance. This ratio is expressed as follows:

$$Q = 2\pi \times \frac{\text{energy stored per cycle}}{\text{energy lost per cycle}}$$

The higher the Q , the greater the amount of energy stored compared with lost energy. However, if this Q is too high, the stored energy produces tank-circuit currents that are too great, and then power is lost in the tank circuit.

If you consider the reason for storing energy in the tank circuit, you will see why an increased Q of a load tank circuit reduces distortion. Energy that is due to generation harmonics is stored in the tank circuit during oscillations in order to enable the flywheel effect of the tank circuit to produce a waveform that is as close as possible to a sine wave. As stated before, plate current flows in a class C oscillator for less than half a cycle, somewhere between 120° and 160° out of the total 360° . During these pulses of plate current, the tube delivers power to the tank circuit, and this power is stored in the tank. When plate current is not flowing, the tank circuit releases this stored energy and produces an a.c. that approaches a sine wave. The greater the energy stored in the tank, the closer is the output to a sine wave. Thus, to produce oscillations that are almost sine waves, it is desirable to store much more energy in the tank than the amount of energy dissipated. This requires a higher Q . Since the harmonics present in the output wave are at the lowest percentage when the output approaches a sine wave, a high

Q of a tank circuit load results in low distortion.

Other factors that contribute to frequency instability are changes in characteristics of tubes, changes in temperature, and vibration. The most important cause of changes in tube characteristics are factors that affect the a.c. plate resistance of the tube. Some of these factors are changes in filament, grid, and plate voltage. The operating voltages can be stabilized by means of a regulated power supply. Another way of counteracting changes in the dynamic or a.c. plate resistance is by using a high value of grid-leak resistance. This increases the grid bias, and raises the effective internal resistance of the tube, thereby making the tube more resistant to changes in plate voltage. The a.c. plate resistance of the tube is also affected by changes in the spacing of the tube elements, which may be due to aging, temperature, or vibration. Since the interelectrode capacitances of the tube depends directly upon the spacing of the elements, any variations in the spacing will affect the frequency of oscillation. In order to minimize these variations, a high L/C ratio may be used. This has the effect of reducing the percentage of change in C and thereby reducing changes in frequency.

In addition to varying the spacing of the tube elements, changes in temperature affect the physical sizes of the tank coil and capacitor, and thus alter the frequency. Temperature effects are much slower than dynamic resistance changes, and frequency changes caused by temperature are known as frequency drift. At higher r-f frequencies, the cause of frequency drift due to changes in room or surrounding temperature is more pronounced. An increase of temperature seems to cause the frequency of oscillators to decrease as if either the L or the C of the tank circuit increased. To overcome this increase of capacitance, a specially designed capacitor that has a negative temperature coefficient can be used. A capacitor that has a negative temperature coefficient is one whose capacitance decreases with an increase in temperature.

30-3. HARTLEY OSCILLATOR

The two oscillator circuits most commonly used in broadcast radio receivers are the tickler-coil and the Hartley. The Hartley is actually a modified version of the tickler-coil oscillator. Instead of using a separate feedback coil and tank coil, the Hartley employs a single coil, tapped in such a way that part of the coil is between plate and cathode, and the other part between grid and cathode. Fig. 30-11a illustrates a series-fed Hartley, and Fig. 30-11b shows a parallel-fed (shunt-fed) Hartley. The difference between the two circuits is in the method of applying d-c plate voltage.

In the series-fed circuit, both the a-c and the d-c components of plate current flow through the L_2 portion of the tank coil. The d-c component flows through the battery and L_2 , whereas the a-c component takes the low-reactance path through C_b and L_2 . In the parallel-fed Hartley, the d-c component of plate current does not flow through L_2 , because its path is from the r-f choke through the battery and then directly to cathode. The a-c component is kept out of the d-c path by means of the r-f choke, which offers a high reactance to r-f currents. The a-c path of the shunt-fed circuit can be traced through the blocking capacitor C_b , to L_2 , and back to cathode. The d-c and a-c paths are entirely separate in a shunt-fed Hartley.

Although both methods of applying plate voltage have certain advantages, the parallel-fed circuit is preferred. The main advantage of the series-fed Hartley is that it does not require an r-f choke. This choke, since it is in parallel with a portion of the tank circuit, may be a source of loss if it does not present a very high impedance at the oscillator's frequency. The choke should have about ten times as much inductance as the tank coil. In addition, the choke has distributed capacitance; therefore, it has a resonant frequency determined by the inductance and the distributed capacities. If this resonant frequency should happen to be the same as that of the tank circuit, the choke would absorb power from the tank.

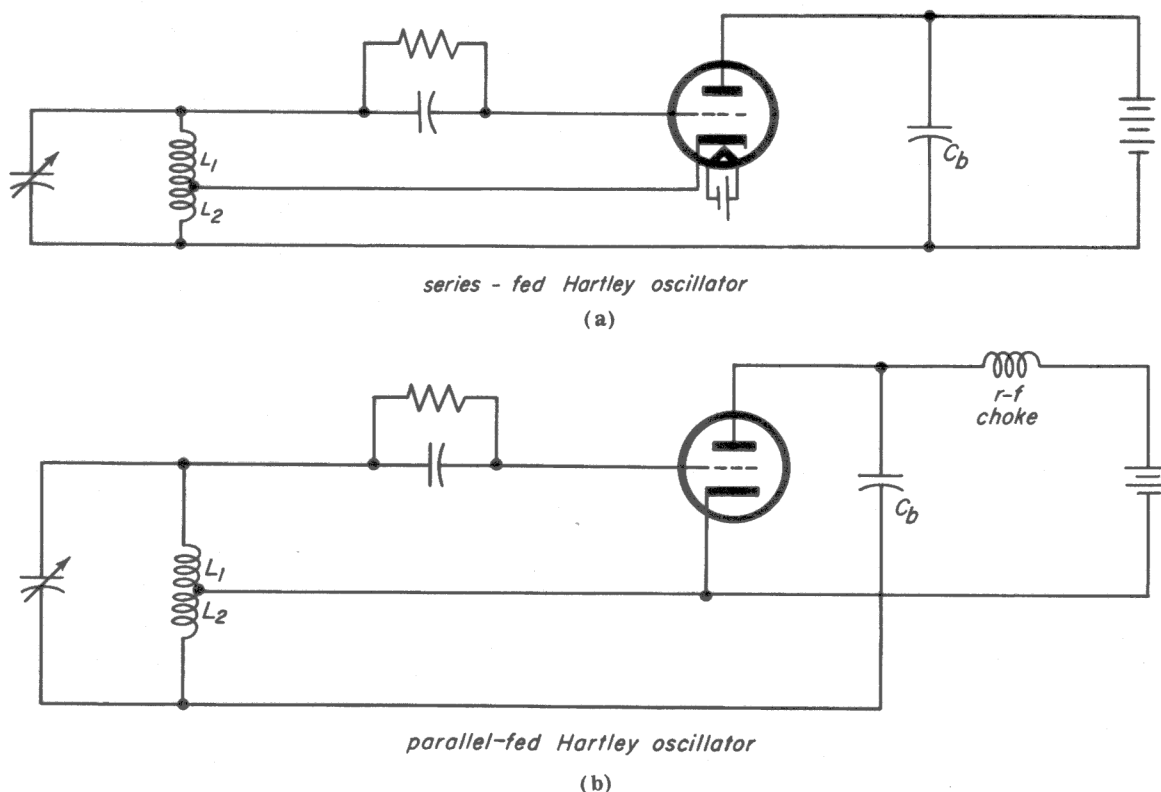


Fig. 30-11

The parallel-fed circuit has the advantage of keeping the d-c voltage out of the tank coil. In this way, the danger of shock is reduced. Secondly, the choke prevents r-f energy from entering the power supply and being coupled to other circuits that use the same power supply. Finally, the tank circuit components in a shunt-fed circuit need be insulated only for r-f currents, and not for high-voltage d.c.

The operation of a Hartley is very similar to that of a tickler-coil oscillator. The lower portion of the tank coil is inductively coupled to the rest of the coil, the combination acting as an autotransformer. The amount of feedback depends on the position of the tap on the tank coil; if the number of turns in the L_2 half of the coil is increased by moving the tap, the feedback is greater, and vice versa. Pulses of plate current flowing through L_2 induce a voltage in L_1 . This induced emf is of the proper phase to sustain oscillations, as in the action of a tickler-coil oscillator. (The two opposite ends of the autotransformer are 180° out of phase. This phase difference, added to the 180° phase shift in the tube, makes a total of 360° , thus bringing the induced voltage in phase with the oscillation). The con-

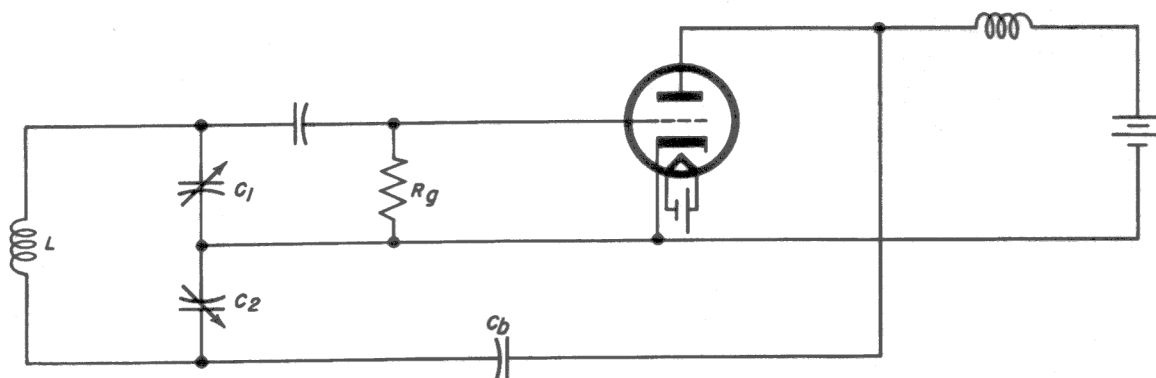
siderations of frequency stability, circuit Q , and loading are the same in a Hartley as in a tickler-coil oscillator.

30-4. COLPITTS OSCILLATOR

The Colpitts oscillator, Fig. 30-12, is similar to the Hartley, except that two capacitors are used with a cathode tap between them, instead of a tapped inductor. Feedback is obtained by means of electrostatic coupling through the ratio of C_1 to C_2 . The circuit is shunt-fed. Grid-leak bias is used, but the grid resistor R_g is connected in parallel with the grid and cathode to provide a d-c path for grid current.

When pulses of plate current flow, C_b starts to charge, and, in turn, plate capacitor C_2 charges. The charge on C_2 is transferred to C_1 through the tank coil. This action continues indefinitely, setting up an oscillating current in the tank circuit. The frequency of these oscillations depends upon the values of L and the equivalent capacitance in the tank circuit.

A pulse of plate current flows for part of the grid-signal cycle, and this energy is



Colpitts oscillator

Fig. 30-12

transferred through C_b to the tank circuit to replace the energy that was lost in the tank circuit and any energy that was withdrawn from the tank circuit by the load.

The two tank capacitors, C_1 and C_2 , act as a voltage divider. The tap between them serves the same purpose as the tap on the inductor in the Hartley circuit, by assuring the proper amount and correct phase of the feedback. When the top of C_1 is negative, the bottom of C_2 is positive, and vice versa. The total voltage across C_1 and C_2 divides in proportion to the reactances of C_1 and C_2 , respectively. This is the same as saying that the voltage divides in *inverse* ratio to the capacitances of C_1 and C_2 (capacitive reactance varies *inversely* with capacitance). Therefore, if you wish to increase the voltage across the grid capacitor C_1 , you must *decrease* the capacitance of C_1 , by making the proper adjustment with the tuning knob. To maintain the same frequency, however, the total tank capacitance must be constant. Therefore, if

you decrease C_1 , you must increase C_2 enough to make the total capacitance equal to the same value as before. The formula for the frequency is, approximately:

$$f = \frac{1}{2\pi\sqrt{LC_t}}$$

where

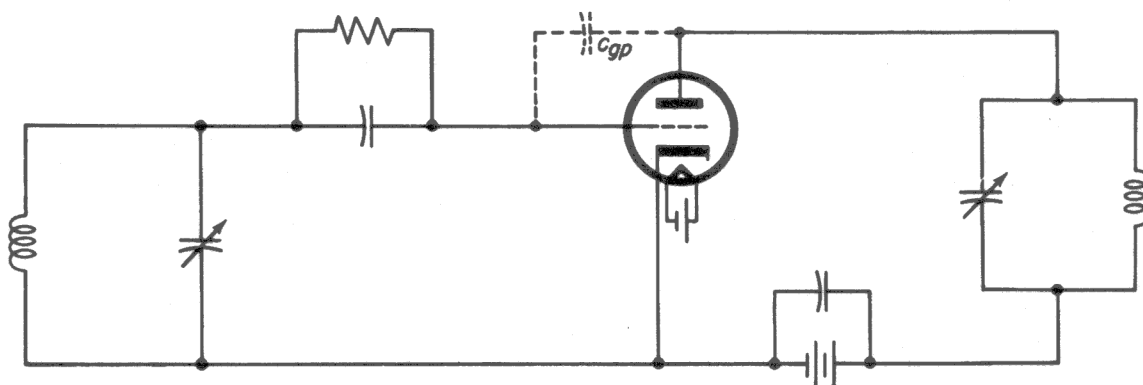
C_t is the total capacitance of the tank circuit.

Since C_1 and C_2 are in series:

$$C_t = \frac{C_1 \times C_2}{C_1 + C_2}$$

30-5. TUNED-PLATE TUNED-GRID OSCILLATOR

This oscillator, abbreviated TPTG, employs a tuned circuit in both the plate and the grid circuits. Figure 30-13 illustrates the TPTG oscillator. Note that the plate circuit



tuned-plate tuned-grid oscillator

Fig. 30-13

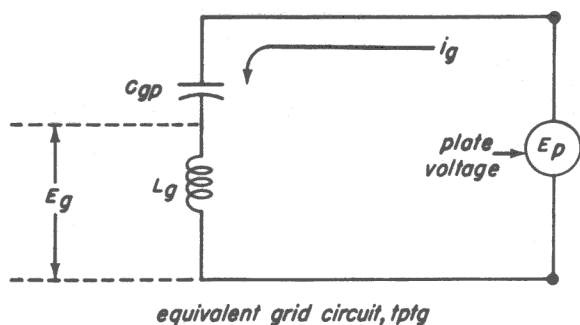


Fig. 30-14

is not inductively coupled to the grid circuit, as was true in the tickler-coil and Hartley oscillators. The feedback necessary for sustaining oscillations occurs through the inter-electrode capacitance between the grid and the plate, shown as C_{gp} in the diagram. The dotted lines signify that C_{gp} is not an actual capacitor.

The TPTG automatically oscillates at a frequency lower than the resonant frequency of both the plate and grid tank circuits. The frequency of oscillation is determined by the tuned circuit that has the higher Q . This may be either the plate tank or the grid tank. The circuit with the higher Q should be tuned closest to the desired frequency. A complete analysis of the phase relations of voltages and currents in the TPTG oscillator is beyond the scope of this course, but a general discussion will help to show that the feedback energy is of the proper phase to keep the circuit oscillating. We know that the plate voltage is 180° out of phase with the grid voltage. The plate voltage is applied across C_{gp} in series with the grid coil (see Fig. 30-14). For the oscillator to operate correctly, the reactance of C_{gp} must be greater than that of the tank coil L_g at the operating frequency. Therefore, the plate voltage causes a current to flow in the grid circuit that is somewhat more than 90° ahead of the plate voltage. (Ordinarily this current would lead the plate voltage by somewhat less than 90° , due to the resistance associated with the circuits. In the TPTG, as in other oscillators, there is negative resistance, accounting for the more than 90° lead mentioned above.) This current in the grid circuit is slightly more than 270° ahead of the grid voltage (180° plus 90°). In flowing through L_g , this current induces the voltage E_g of the grid tank circuit. The voltage across L_g must lead

the current by somewhat less than 90° since there is some resistance. This makes the induced voltage E_g in phase with the grid voltage (270° plus 90° equals 360°).

30-6. CRYSTAL OSCILLATOR

If we remove the grid tank circuit from a TPTG oscillator and replace it with a crystal, we have a crystal oscillator. The control of frequency by means of crystals is based on the piezoelectric effect, discussed in previous lessons. When certain crystals are compressed or stretched in specific directions electric charges appear on the surface of the crystals. Also, such crystals expand or contract when placed between two metallic surfaces across which there is a difference of potential. If a slice of crystal is compressed so that it bulges inward as in Fig. 30-15a, opposite charges appear on its faces. If the same crystal is squeezed in such a way that it bulges outward, the charges are reversed (Fig. 30-15b). If the crystal is made to bulge inward, and outward at regular intervals, it becomes a source of alternating voltage. When a source of a.c. is applied to a crystal, it vibrates back and forth at its natural frequency, which depends upon its structure. The amplitude of the crystal vibrations will be greatest when the frequency of the a.c.

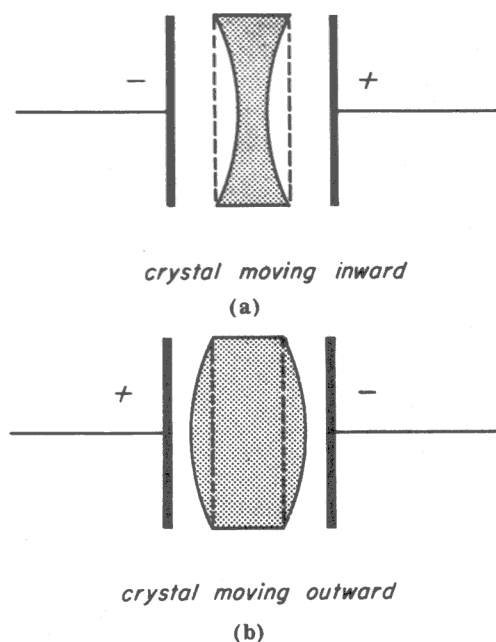


Fig. 30-15

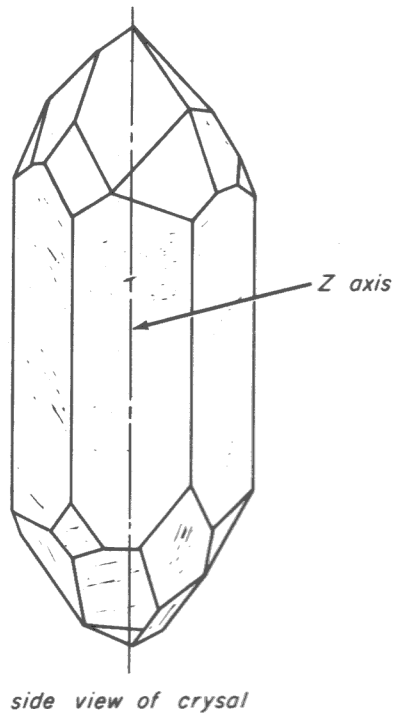


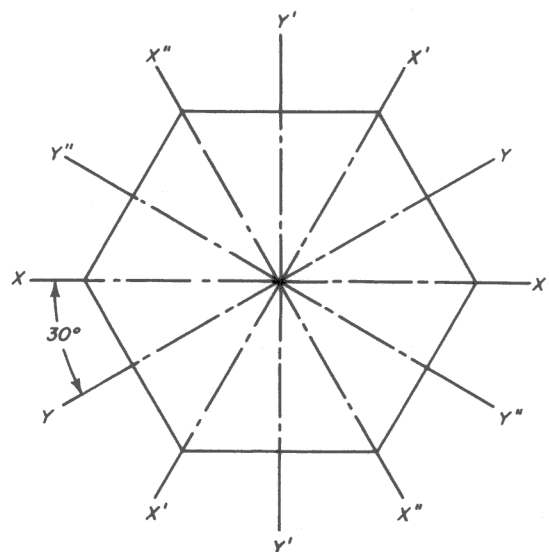
Fig. 30-16

voltage is equal to the resonant frequency of the crystal. The vibrations at the crystal's natural frequency will sustain themselves and generate electrical oscillations if all mechanical losses are overcome. Thus, a crystal can be substituted for the tank circuit in an oscillator. The chief advantage of using a crystal is that it can maintain the output frequency of an oscillator at a constant value for long periods of time. Before we discuss the crystal oscillator, let us find out a little more about crystals.

The piezoelectric effect is present in all types of crystals, but the ones that are considered for use in electric circuits are tourmaline, Rochelle salt, and quartz. Tourmaline is expensive because it is a semi-precious stone. It is used sometimes in high-frequency oscillators, because it need not be made as thin as quartz would have to be at these frequencies. Rochelle salt can be made to generate the greatest amount of voltage for a given mechanical strain, but it is not suitable for frequency control because it is unstable. Rochelle salt is used in microphones, crystal speakers, and phonograph pick-ups. Quartz is the best of the three for frequency control, because it is inexpensive, mechanically rugged, and expands very little with heat.

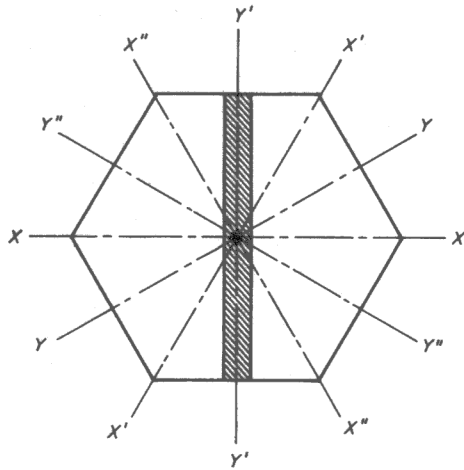
Natural quartz crystals are always shaped like a hexagon (a figure with six sides), with ends shaped like a pyramid (see Fig. 30-16). When the crystal is mined, one or both of the points is usually broken off. The dotted line joining the two points in the diagram is called the optical or Z axis. This axis shows no electrical effects; that is, when the crystal is compressed along this axis, no differential of potential is produced. Fig. 30-17 is a cross section view of a crystal, with the X and Y axes shown. The three X axes, X, X', and X'', pass through the corners of the cross section; these axes are the *electrical* axes, because they are the directions along which the greatest piezoelectric effect is produced. The Y axes, Y, Y', and Y'', are known as the *mechanical* axes, and they pass through the faces of the crystal at right angles. Both the X and Y axes are perpendicular to the Z axis, which runs lengthwise through the center of the crystal, from point to point.

When a voltage is applied to any X axis, it causes an expansion or a contraction along the Y axis that is at right angles to the particular X axis. For example, a voltage applied to the X axis of Fig. 30-17 will produce a mechanical strain along the Y' axis. If a crystal is compressed along any Y axis, a voltage will appear on the faces of the crystal along the X axis that is at right angles; thus, squeezing of the crystal along the Y'' axis will cause a voltage to appear along the X' axis.



cross section of crystal

Fig. 30-17

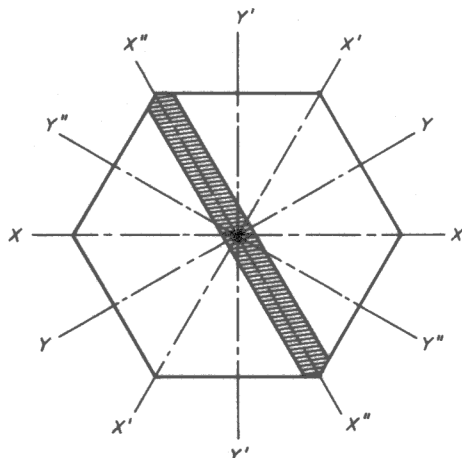


X-cut

Fig. 30-18

There are many ways that slabs or slices of crystal can be cut from the natural crystal. The different ways of cutting slices result in a variety of crystal *cuts*, which are identified by certain letters. Each cut has a different set of characteristics. In general, the most desirable cuts will have a single frequency of oscillation, ease of oscillation at the intended frequency, and a minimum of frequency changes due to temperature changes.

A cut that is sliced along a Y axis and has its parallel faces at right angles to an X axis is called an *X-cut* (see Fig. 30-18). If the crystal is cut along an X axis, and its parallel faces are perpendicular to a Y axis, it is known as a *Y-cut*, or *30°-cut* (see Fig. 30-19). The name *30° cut* comes from the fact that Y axis that goes through the center of the Y cut is at an angle of



Y-cut

Fig. 30-19

30° with respect to the nearest X axis. The X cut and the Y cut are older cuts, and are not used often because they are affected by temperature changes, which cause variations in the frequency of oscillation. However, oscillators using crystals are more stable than the types previously discussed.

A number of cuts have been developed that overcome the disadvantages of two cuts mentioned above. If a slice is made by rotating about the X axis so that the angle is 35° to the right of the Z axis, we have an *AT cut*. If the angle is 49° to the left of the Z axis, the result is a *BT cut*. These two cuts are suitable for operation above 500 kc. They show *no* frequency variation when the temperature is between 25° and 45° C. This is known as having a zero temperature coefficient between those two temperatures. Other crystal cuts, with their angles of cutting, and useful frequencies are: *CT cut*, 38° to right of Z axis, between 50 and 500 kc; *DT cut*, 52° to left of Z axis, between 50 and 500 kc; *ET cut*, 66° to right of Z axis, 100 to 1,000 kc; *FT cut*, 57° to left of Z axis, 100 to 1,000 kc; *GT cut*, obtained by rotating the principal axes of the *CT* or *DT* cut by 45°, 100 to 500 kc.

When crystals are used in electronic circuits, they are placed within a holder, between two metallic electrodes. A capacitor is thus formed, the crystal itself serving as the dielectric. The holder should not interfere with the vibrations, but should hold the crystal firmly in position. The clamp-type of holder actually clamps the crystal in place between the two electrodes; the air-gap type of holder uses an air gap between the crystal and one or both of the electrodes. An adjustable air gap can be used to permit slight adjustments in the frequency.

A crystal acts like a tuned circuit at its resonant frequency. The equivalent electrical circuit is shown in Fig. 30-20. C_m is the capacitance of mounting; C_g is the effective series capacitance introduced by the air gap when the electrodes do not touch the crystal. The series combination of L , R , and C forms a resonant circuit that represents the electrical equivalent of the crystal's mechanical resonance. The frequen-

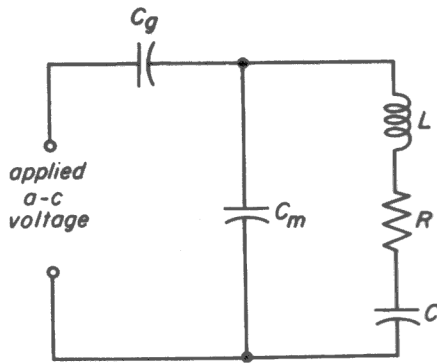


Fig. 30-20

cy at which L and C are in series resonance is the frequency of mechanical crystal resonance. C_m is in parallel with the L , C , and R combination, and, therefore, the circuit has a parallel-resonant frequency also slightly above series resonance. The parallel resonant frequency occurs when the inductive reactance of the series branch is equal to the capacitive reactance of C_m . The reactance of C_m is comparatively low, and thus only a small inductive reactance is required in the LCR branch to produce parallel resonance. This makes the series- and parallel-resonant frequencies very close together. The frequency-current characteristic is shown in Fig. 30-21.

Tuning is accomplished by varying the frequency from a value below crystal resonance to one above crystal resonance. When series resonance is reached, current is maximum (see the high peak in the curve of Fig. 30-21). In spite of the large current, oscillations cannot start because of improper phase relationships. When parallel resonance is reached, current drops to minimum, because this is the point of highest

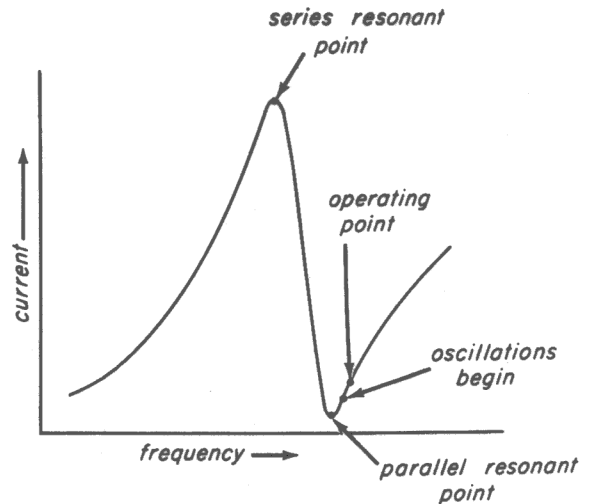


Fig. 30-21

crystal impedance. A condition of oscillation in a crystal oscillator, as in a *TPTG*, is that the plate circuit be *inductive*. At parallel resonance, the circuit is resistive; therefore, oscillations do not start at this point. When the tuning brings the frequency to a point slightly above parallel resonance, the plate circuit becomes inductive, and oscillations begin. The tuning should be adjusted to a frequency slightly above that at which oscillations begin, in order to insure stable operation over long periods of time.

A typical crystal oscillator circuit is shown in Fig. 30-22. As in the *TPTG*, feedback is obtained through C_{gp} . The choke in series with R_g prevents r-f energy from entering the d-c grid-current path. The crystal stores energy in mechanical form during half of the cycle and releases it in electrical form during the other half-cycle. The rate of storage and release of energy is de-

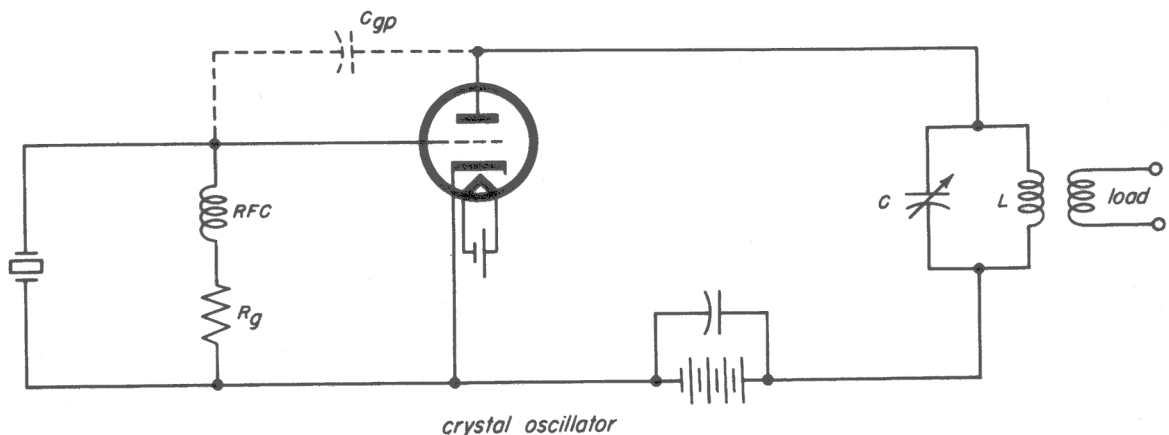


Fig. 30-22

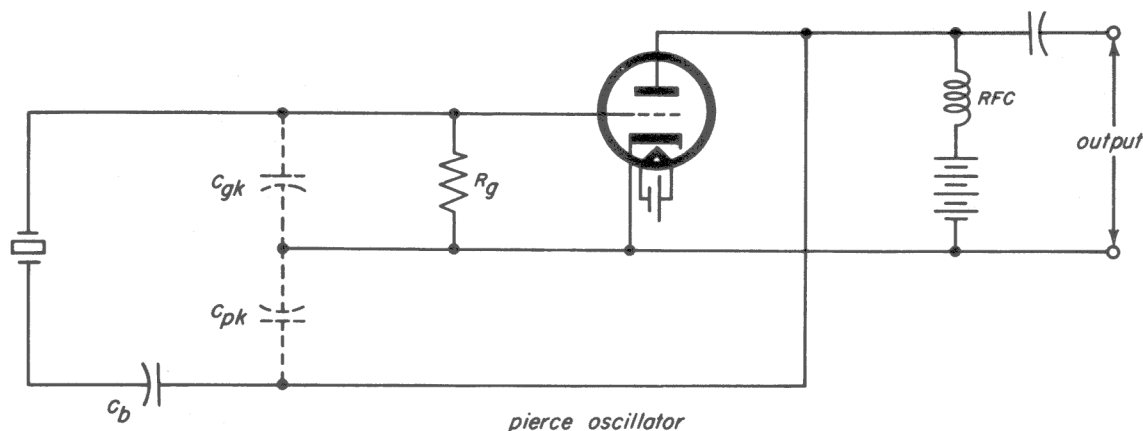


Fig. 30-23

terminated by the natural resonant frequency of the crystal. The energy losses are overcome by the feedback occurring through C_{gp} .

Fig. 30-23 illustrates a Pierce oscillator. Inspection of this circuit will show that it is the same as a Colpitts, with the crystal taking the place of the tank coil L of Fig. 30-12, and capacitors C_1 and C_2 being replaced by C_{gk} and C_{pk} . The circuit capacitances of a Pierce oscillator do not affect the frequency of operation, because the crystal itself serves as the tuned circuit. C_{gk} , the grid-to-cathode capacitance of the tube, and C_{pk} , the plate-to-cathode capacitance, act as a voltage divider, and determine the amount of feedback from plate to grid. The Pierce is useful in high-frequency oscillators that involve low power, such as oscillators used for testing and calibration purposes.

30-7. MISCELLANEOUS OSCILLATORS

Two types of oscillators are commonly used for reception of Morse code signals, when transmitted as *CW* (continuous wave) signals. In one type, a special detector circuit, known as a *regenerative* detector, is used. As the term suggests, a regenerative detector operates on the principle that positive feedback will reinforce the incoming signal, and produce a greatly amplified signal in the output. The feedback must be kept within limits in this detector. When the feedback is increased to a certain point, the circuit will break into oscillation, and it becomes an *oscillating detector*. The oscillations thus produced are made to beat

(heterodyne) with the incoming signal. The output contains an audio signal, whose frequency is the difference between the incoming signal frequency and that of the oscillating detector. For example, an incoming 3,000-kc signal will beat with a 3,001-kc oscillation to produce a 1-kc (1,000-cycle) audio signal.

A beat-frequency oscillator (*BFO*) is also used to receive *CW* signals. The *BFO* is found as an additional oscillator in a superheterodyne receiver, to produce a frequency that will beat with the i.f. and establish a difference frequency in the audio range.

Other oscillators that have special applications in high- and ultra-high frequency circuits are beyond the scope of this course. Some of these oscillators are used in radar equipment. These are the names of some of these oscillators: Barkhausen-Kurz, klystron, magnetron, Wien-bridge, thyatron sawtooth, RC oscillator, relaxation, multi-vibrator, gas-diode, and blocking oscillators.

30-8. COUPLING

There are several ways in which the oscillator voltage may be transferred to another circuit. One method is by inductive coupling, as shown in Fig. 30-24a. In this type of coupling, a second coil is placed near the tank's coil. This coupling is similar to that in a transformer; the second coil can be considered as a secondary. We know that there is an a-c current flowing in the tank circuit. Also, that there is an expanding and

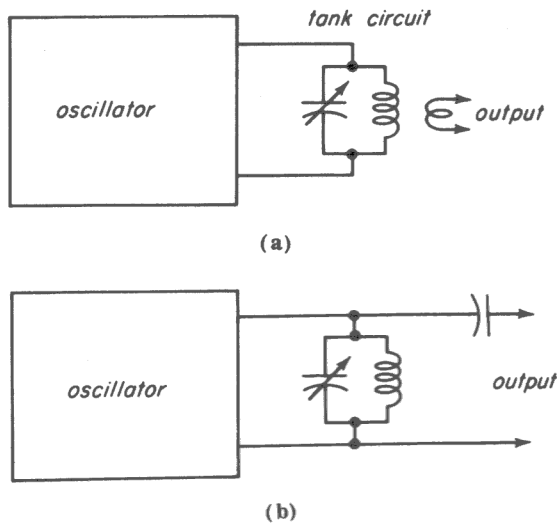


Fig. 30-24

contracting magnetic field about the coil. When the field cuts the turns of the secondary coil, a voltage is induced in the secondary. This voltage can be applied to another circuit. The magnitude of the voltage induced in the secondary depends upon the number of turns in both coils and amount of coupling between the two coils. The tighter the coupling, the greater will be the voltage induced in the secondary.

Another type of coupling is capacitive coupling, shown in Fig. 30-24b. This type of coupling is simple. Since there is an a-c voltage across the tank circuit, this voltage can be applied to another circuit by connecting one end of the circuit to the ground end of the tank coil and the other end to the grid end, through a capacitor. The voltage of the tank divides between the reactance of the coupling capacitor and the resistance of the next circuit. A large capacitance (low reactance) will result in a great proportion of the available voltage being applied to the following circuit. In some cases, the capacitor may be variable so that the amount of coupling can be controlled.

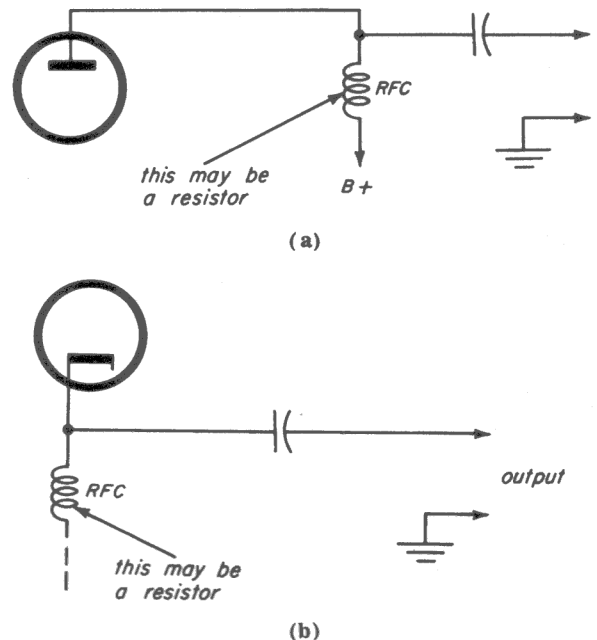


Fig. 30-25

The coupling circuits shown in Fig. 30-24 can be used with any class of tube operation. Another basic type coupling is also used in circuits using class A_2 oscillators. This circuit is shown in Fig. 30-25a and b. Since the plate current flows in class A_2 oscillators for 360° of the input signal cycle, the plate (and cathode) current will be relatively undistorted. A coil (called a radio-frequency choke) or a resistor connected in the cathode or plate circuit will have developed across it an r-f voltage of the same waveform. The coupling to the next stage can be either by an impedance or a resistance-capacitance coupling as shown in the illustration. Figure 30-25a shows the RFC in the plate circuit, and Fig. 30-25b shows the RFC in the cathode circuit. This type of coupling may be used in class C operation if harmonics (due to the distorted plate current) are attenuated by tuned circuits that will accept only the fundamental-frequency signal.